

# Enhancing Identity Documents Security: Developing a Verification Tool to Prevent Spoofing, Tampering, and Forgery

Mamoon Obiedat<sup>1</sup>, Musab Al-Ghadi<sup>1\*</sup>, Ashraf H. Aljammal<sup>2</sup>, Rawa Alzghaybh<sup>1</sup> and Raghad Abu Wadi<sup>1</sup>

<sup>1</sup>Department of Information Technology, Faculty of Prince Al-Hussein Bin Abdallah II for Information Technology, The Hashemite University, Zarqa, Jordan.

<sup>2</sup>Professor, Department of Computer Science and Applications, Prince Al Hussein bin Abdullah II Faculty of Information Technology, The Hashemite University, Zarqa, Jordan.

\*Corresponding author: Musab Al-Ghadi. Email: [alghadi@hu.edu.jo](mailto:alghadi@hu.edu.jo)

Received: March 10, 2026 Accepted: May 20, 2026

**Abstract:** Remote onboarding is now vulnerable to AI-created deepfakes and identity fraud, particularly given the increased compliance standards introduced by eIDAS 2.0, which mandates highly robust identity document verification as a crucial component of secure online operations. This paper proposes a multimodal deep learning model that consists of three unified modules designed to overcome these problems. The initial module is a visual similarity analysis to ensure identity with the help of a Siamese Convolutional Neural Network. The second module applies specific fraud detection in different countries with dynamic thresholds to consider the differences in documents and attacks between jurisdictions. The third module separates original and scanned documents with the help of an EfficientNet-B7-based classifier. The system was trained and tested using the MIDV-2020 dataset that consists of 10 country documents (320 real and 320 forged documents per country). Experimental measures show a True Acceptance rate of 99-100%, a False Acceptance rate of 6.0%, and a pipeline latency of 278 ms on CPU and 42 ms on GPU, which can be used to run GUI operation in real time. The proposed multimodal verification framework integrates structural verification, adaptive fraud detection, and presentation attack detection into a single pipeline, ensuring more robust and reliable digital identity verification while meeting regulatory requirements. The next step in work will be INT8 quantization to make it possible to use it on mobile and IoT platforms efficiently.

**Keywords:** Identity document verification; identity fraud; remote onboarding; visual similarity; deepfakes; document fraud, deep learning; MIDV-2020

## 1. Introduction

This is bolstered by the high pace of the digital service and remote transaction growth that has led to the increased demand for secure, reliable, and user-friendly digital identity verification systems and digital trust that is enforced by regulatory forces to protect individual data. The remote onboarding has been accelerated by the COVID-19 pandemic, which necessitates the use of online identity verification in banking, fintech, e-government, and telemedicine. Simultaneously, the threat environment has also changed with the introduction of AI-driven deepfakes and sophisticated graphic editing tools that can create highly realistic forged documents that can be difficult to detect with automated checking systems [1-3]. This also demonstrates the value of proper identification of citizens [4] and has resulted in increased legal safeguards of documents and digital identities [5]. The advanced security features incorporated into modern identity documents, including guilloche patterns [6,7], holograms [8-12], anti-scan patterns, watermarks, and micro-printing, are embedded into the identity document. Some examples of international standards that provide machine-readable document specifications are ICAO Doc 9303 [3], whereas standards like NIST SP 800-63 [13,14] define the level of identity assurance, and projects like

eIDAS 2.0 encourage privacy-conscious digital identity ecosystems. Irrespective of these developments, most verification systems are still limited, most of them based on single-modality examination like facial biometrics or isolated feature examination, hence prone to advanced fraud. There is also a uniform decision threshold; cross-country robustness is also weakened. Even though the existing literature discusses visual similarity learning [15,16], fraud detection pipelines [17], and texture-based scanned document detection [18,19], a document-centric approach has not been available. There exists a necessity of a single, regulation-congruent architecture that can concurrently assess structural similarity, implement country-specific adaptive calibration, and identify scanned or digitally reconstructed presentation attacks on-the-fly. These gaps are essential in resolving strategies to attain an improved level of identity assurance during remote onboarding without compromising the efficiency of its operations and regulatory suitability.

To overcome these shortcomings, this paper proposes a single deep learning-based identity document verification system with the following contributions:

- i. Integrated multi-module architecture: This refers to a system in which structural similarity analysis, adaptive fraud detection, and presentation attack detection are performed in one system.
- ii. Siamese CNN-based visual similarity analysis: A similarity-learning network is a network that compares identity documents with certified templates to detect structural manipulation and visual forgery [15,16].
- iii. Adaptive thresholds in country-specific fraud detection: A country-sensitive calibration mechanism instead of general thresholds is used to enhance the ability to discriminate between authentic and fabricated documents in a variety of forms [17].
- iv. Original-versus-scanned classification with an EfficientNet: The paper trains an EfficientNet encoder using a multilayer perceptron classifier, which identifies documents scanned or reconstructed digitally by detecting texture and compression artifacts, and is consistent with the presentation attack detection principles of ISO/IEC 30107-3 [18-20].

The combination of all these complementary elements helps the framework enhance resilience to complex document fraud, facilitates greater levels of identity assurance, and meets international standards to ensure the safe deployment of the framework in real-time in banking, governmental, and border control settings [21,22].

The rest of this paper is structured in the following way. Section 2 provides a literature review regarding related research in the areas of identity document verification and presentation attack detection. Section 3 will present the suggested multi-module identity verification system and its methodological aspects. Section 4 details the dataset, experiment design, and the evaluation measures. The results of the experiment are provided and discussed in Section 5. The comparative study and the discussion are presented in section 6. Section 7 has the prototype user interface design and system workflow. Lastly, the paper ends with the future research directions as optimization of models to mobile and IoT environments.

## 2. Related Works

Fraud in documents has been more advanced and made more accessible [23]. Due to the increase in digital commerce and remote onboarding, fraud prevention systems will need to change. By combining Fraud Prevention Systems (FPS) with Fraud Detection Systems (FDS), it is possible to reinforce the electronic commerce infrastructures [16], and approaches based on databases like BIDIF, which investigate graphical features of the forged documents [5].

### 2.1. Taxonomy of Existing Identity Verification Approaches

The existing identity verification systems are classified into various broad categories. Facial recognition solutions are concerned with biometric similarity between document images and live captures [24,25], which offer high identity verification and low document verification. Structural conformity analyzes layout and character-level properties to identify manipulations such as scan-edit-print [17,26], but cannot do so with global layout. Cheating Perceptual and visual consistency techniques evaluate similarity and layout consistency with high-quality templates [15,21,27,28], though they cannot be applied effectively to a wide range of document types. Content consistency checking evaluates the textual and semantic consistency [27] [31-32], which is based on proper OCR and organized data extraction. The quality of images and the liveness-based analysis of presentations identify presentation attacks and distinguish between original and scanned or copied documents [18,19]. Lastly, region-based models divide documents

into key areas to perform localized authenticity verification [8,9,29], improving their interpretability without a strict localization constraint and document-specific suppositions.

## 2.2. Analytical Comparison and Identified Gaps

Although prior research provides important contributions across biometric, structural, perceptual, and quality-based verification dimensions [9,15,17-19,24-29], several limitations remain. Many existing systems rely on single-modality analysis, which reduces robustness against complex fraud scenarios combining structural preservation, visual manipulation, and reproduction attacks. Structural methods [17,26] can detect layout and character inconsistencies but may fail when global layout conformity is preserved. Perceptual similarity approaches [15,21,27,28] are effective against visual tampering, yet often lack adaptive, country-specific calibration, limiting generalization across heterogeneous document templates. Quality and liveness-based techniques [18,19] strengthen the detection of scan or reproduction attacks but are frequently implemented as isolated components without integration into broader structural or semantic verification pipelines. Likewise, region-based detection frameworks [8,9,29] enable localized analysis of critical security features, yet often depend on document-specific configurations, reducing scalability. Overall, unified frameworks that combine visual similarity, adaptive thresholding, presentation attack detection, and region-level analysis within a single scalable architecture remain limited. Developing a robust, multi-layered verification pipeline capable of handling diverse document layouts and evolving fraud techniques, therefore, remains an open research challenge. Table 1 presents an analytical summary of the related works.

**Table 1.** Analytical Summary of Related Works.

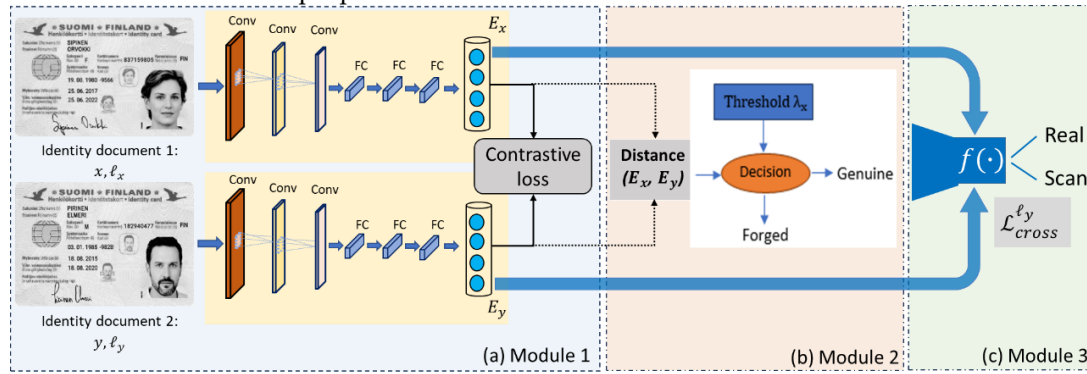
| Related work | Main Approach                                       | Strengths                                     | Limitations                                 | Identified Gap                              |
|--------------|-----------------------------------------------------|-----------------------------------------------|---------------------------------------------|---------------------------------------------|
| [24,25]      | Facial recognition                                  | Strong biometric validation                   | Does not verify document authenticity       | Lacks document-centric fraud analysis       |
| [26]         | Character-level structural analysis (SEP detection) | Detects scan-edit-print manipulation          | Sensitive to subtle structural preservation | Limited generalization                      |
| [17]         | Structural conformity                               | Layout compliance verification                | Uniform thresholds across formats           | No adaptive country calibration             |
| [15,28]      | Visual similarity learning                          | Detects perceptual inconsistencies            | Template dependency                         | Limited integration with liveness detection |
| [21]         | Perceptual consistency                              | Robust visual matching                        | May require extensive training data         | Lacks presentation attack detection         |
| [27]         | Content consistency                                 | Semantic validation                           | OCR dependency                              | Vulnerable to image degradation             |
| [18]         | FFT-based scan detection (CheckScan)                | Effective original vs. scanned discrimination | Focused on frequency domain only            | Not integrated with structural analysis     |
| [19]         | DLD layout descriptor                               | Local and global spatial matching             | Layout-dependent                            | Limited multimodal integration              |

|          |                                         |                                            |                       |         |                                    |
|----------|-----------------------------------------|--------------------------------------------|-----------------------|---------|------------------------------------|
| [8,9,29] | Region-based security feature detection | Targeted verification of security elements | Requires localization | precise | Lack of unified adaptive framework |
|----------|-----------------------------------------|--------------------------------------------|-----------------------|---------|------------------------------------|

As can be seen in Table 1, the literature has no integrated and adaptive framework that integrates structural verification, country-adaptive fraud thresholds, scanned documents detection, and live deployment. This gap is important to fill in the digital onboarding and highly secure environments in terms of robust identity assurance.

### 3. Proposed Multi-Module Identity Verification Framework

To overcome the shortcomings of single-modality verification systems, this section introduces a multi-module country-adaptive identity verification system. The structural similarity analysis, real-time adaptive fraud decision-making, and presentation attack detection modules are integrated into one graphical display. The framework consists of three complementary modules: (i) Similarity-Based Structural Verification, (ii) Country-Specific Fraud Detection, and (iii) Original vs. Scanned Document Classification. Unlike conventional single-task systems, the proposed architecture combines structural similarity learning, adaptive decision thresholds, and presentation attack detection within a single deployment pipeline. Figure 1 presents the framework of the proposed modules.



**Figure 1.** The framework of the three proposed modules. (a) similarity-based structural verification module. (b) country-specific fraud detection module. (c) original vs. scanned document classification module.

#### 3.1. Module 1: Similarity-Based Structural Verification

This module measures the structural and visual similarity of a query document and a reference template with the help of a Siamese Convolutional Neural Network (Siamese CNN). The network is trained to learn a common embedding space in which visually similar documents are brought near to one another. Each Siamese branch consists of three convolutional blocks (reflection padding, Conv2D, ReLU, batch normalization) followed by fully connected layers that reduce the flattened feature space (720,000 features) to a compact 5-dimensional embedding vector. Both branches share weights to ensure consistent feature extraction. Input images are converted to grayscale, resized to  $300 \times 300$  pixels, normalized, and passed through the network. Structural similarity is measured using the Euclidean distance:

$$d = \|E_x - E_y\|_2$$

Lower distance values indicate higher structural similarity.

This proposed module focuses purely on structural and visual similarity without making a final authenticity decision. Fig 1(a) presents the framework of the proposed module 1.

##### 3.1.1. Network Structure

The Siamese network consists of two identical subnetworks that use the same weight, denoted forward once, which obtains deep features of both the template and query image. Each subnetwork comprises a lightweight convolutional backbone (CNN1) consisting of three convolutional blocks. Each block applies 1-pixel reflection padding, followed by a  $3 \times 3$  Conv2D layer with progressive channel expansion ( $1 \rightarrow 4$ ,  $4 \rightarrow 8$ , and then  $8 \rightarrow 8$ ), kernel size = 3, batch normalization, and ReLU activation. The stride is set to 1 to preserve spatial resolution, enabling effective extraction of fine-grained structural features. These layers are designed to extract local features from grayscale image inputs.

The flattened feature map of  $8 \times 300 \times 300$  is fed to the fully connected module (fc1), to produce a feature vector of 720,000 dimensions. Three linear layers of ReLU activation (Linear layer  $720,000 > 500$ ) followed by ReLU, another Linear layer ( $500 > 500$ ) with ReLU activation, and a final Linear layer ( $500 > 5$ ) generate a 5-dimensional embedding. This architecture gradually reduces the high-dimensional image representation into a low-dimensional feature space, which allows a high degree of similarity comparison of input images.

### 3.1.2. Input Preprocessing and Flow

The input images (template and query) are processed into grayscale, resized 300x 300 pixels, turned into PyTorch tensors, and normalized and batched (unsqueeze). These processed images are fed through every one of the twin subnetworks to produce embedding:

*output1, output2 = model(test\_tensor, forgery\_tensor)*

then Euclidean distance between two output vectors is computed:

*distance = F.pairwise\_distance(output1, output2)*

This is the similarity score that is a scalar distance. A lower value means high similarity (i.e., then it is probably a match) and a high value means dissimilarity. The pseudo-code of Siamese CNN Similarity Calculation (module 1) is given in the algorithm 1.

---

#### Algorithm 1: Siamese CNN Similarity Calculation

---

**Input:** *template\_image, query\_image*

**Output:** *similarity\_distance*

**Step 1.** *preprocess(template\_image) → grayscale, resize(300×300), normalize*

**Step 2.** *preprocess(query\_image) → grayscale, resize(300×300), normalize*

**Step 3.** *embedding1 ← CNN\_forward(template\_image)*

*# 3 Conv blocks (1→4→8→8ch) + FC (720K→500→500→5D)*

**Step 4.** *embedding2 ← CNN\_forward(query\_image)*

*# Shared weights twin network*

**Step 5.** *distance ← Euclidean(embedding1, embedding2)*

**Step 6.** *return distance*

*# Low value = high similarity*

---

Algorithm 1 measures structural similarity between a template and a query document using a Siamese CNN. The images are preprocessed (grayscale, rescaled to 300 x 300, normalized), and each of them is run through the same CNN branches to generate 5-dimensional embeddings. The similarity score is the Euclidean distance between embeddings: a smaller value means that the two are similar, and a larger one means that there is a possibility of forgery. The algorithm, therefore, reduces high-dimensional images to compact representations that are used in verifying structural authenticity.

### 3.1.3. Model Deployment and Execution Flow

The model weights are dynamically loaded (id model.pt), which enables them to support various kinds of documents or various regions. Inference is done with the help of *torch.no\_grad()* so as to minimize memory and speed up the process. The user interface is given the final similarity score.

## 3.2. Module 2: Country-Specific Fraud Detection

This module goes further and uses structural similarity analysis to do decision-level fraud classification with country-specific thresholds that are adaptive. Based on a query image and a chosen label of a country, the system retrieves an authentic template and calculates the similarity distance with the help of the Siamese CNN. That outcome is contrasted with a calibrated threshold  $\lambda_c$ , which is established by cross-validation of authentic and fake samples of each country.

Decision rule:

- If  $d < \lambda_c \rightarrow$  Genuine
- Otherwise  $\rightarrow$  Forged

The result and duration of execution are shown in the GUI. This component, as opposed to module 1, which produces a similarity measure, presents country-sensitive adaptive thresholding, which makes cross-national deployment scalable and allows better robustness on heterogeneous document layouts. The

structure of the proposed module 2 is shown in Fig 1(b), whereas the pseudo-code of the country-specific fraud detection module is shown in Algorithm 2.

---

**Algorithm 2: Country-Specific Fraud Detection**


---

**Input:** *query\_image, country\_label*  
**Output:** *verdict ("Genuine" | "Forged")*  
 Step 1. *reference*  $\leftarrow$  *select\_random\_template(country\_label)*  
 Step 2. *preprocess(query\_image, reference)*  $\rightarrow$  *grayscale, 300×300*  
 Step 3. *distance*  $\leftarrow$  *Siamese\_CNN(query\_image, reference)*  
       *# Call Algorithm 1*  
 Step 4. *threshold*  $\leftarrow$  *load\_country\_threshold(country\_label)*  
       *#  $\lambda_c$  calibrated*  
 Step 5. *if distance < threshold:*  
       *return "Genuine"                      # Green GUI indicator*  
 Step 6. *else:*  
       *return "Forged"                      # Red GUI indicator*

---

The country-adaptive thresholds are used by Algorithm 2 to determine whether an identity document is genuine or forged. The system loads a reference template of a chosen nation, preprocesses query and reference pictures (grayscale, 300 × 300 pixels, normalized) and forwards them to the Siamese CNN (Algorithm 1) to calculate the Euclidean distance of embeddings. The classification is based on a country-specific threshold (  $\lambda_c$  ) in which a distance below  $\lambda_c$  is called Genuine, and otherwise Forged. The GUI displays the results, which allow the detection of fraud to be available in a range of document formats.

### 3.3. Module 3: Original vs. Scanned Document Classification

The module suggested helps decide whether an identity document image is an original capture or a scanned copy, which is complemented by the structural verification stage. This is significant in forensic identity analysis, since scanned documents frequently reflect reproduction, editing, or malicious misuse and generally display compression artifacts, diminished sharpness of sides and a constant background noise pattern. The pipeline of detection has a two-stage deep learning design. The first step is that an EfficientNet-B7 encoder, which is pre-trained with the ImageNet statistics, is used to extract high-level semantic and texture features on a 224 × 224 RGB image, resulting in a 1000-dimensional feature vector. Second, an MLP classifier (two hidden layers of 128 and 32 neurons with the ReLU activation) is trained to predict this feature vector with a binary outcome (original vs. scanned) by means of a SoftMax output layer. The end-to-end training of the model is based on categorical cross-entropy loss on a balanced set of original documents and scanned documents under different conditions of capture, to improve the generalization. During inference, the model parameters are dynamically loaded and stored in country-specific checkpoint files. This system takes the input image and preprocesses them, followed by feature extraction, and the result is a binary label (0 = original, 1 = scanned), which defaults to a scanned label if there is inference failure. This module is incorporated into the GUI in its entirety, and results and processing time are shown in the form of intuitive visual indicators. This proposed module operates independently of structural similarity and focuses on forensic texture and compression artifacts, strengthening the detection of scan-based and reproduction attacks. Fig 1(c) presents the architecture and inference process of the proposed module 3. Algorithm 3 presents the pseudo-code of the proposed module 3 to differentiate between original and scanned documents.

---

**Algorithm 3: Original vs. Scanned Detection**


---

**Input:** *input\_image, country\_label*  
**Output:** *classification ("Original" | "Scanned")*  
 Step 1. *preprocess(input\_image)*  $\rightarrow$  *resize(224×224), RGB, ImageNet normalize*  
 Step 2. *features*  $\leftarrow$  *EfficientNet\_B7(input\_image)*  
       *# Pre-trained, 1000D features*  
 Step 3. *logits*  $\leftarrow$  *MLP(features)*  
       *# 1000→128→32→2 classes*  
 Step 4. *probabilities*  $\leftarrow$  *SoftMax(logits)*

---

---

```

Step 5. if  $\text{argmax}(\text{probabilities}) == 0$ :
    return "Original"    # Green GUI indicator
else:
    return "Scanned"    # Red GUI indicator

```

---

Algorithm 3 classifies an identity document image as either an original capture or a scanned/digitally reproduced copy using a deep learning-based presentation attack detection model. Initially, the input image is preprocessed by resizing it to  $224 \times 224$  pixels, converting it to RGB format, and normalizing it using ImageNet statistics to meet the requirements of the pre-trained encoder. The resulting processed image is subsequently delivered to an EfficientNet-B7 encoder, which attains a 1000-dimensional feature representation of high-level semantic and texture appearance. These characteristics are then input into a multilayer perception (MLP) classifier of the following architecture 1000-128-32-2 to produce logits of the two classifications: Original and Scanned. A SoftMax function transforms the logits into probabilities of the classes, and the most probable of the classes is chosen as the final prediction. Given that the index of the predicted class is 0, the document will be labeled Original, and otherwise, scanned with the result being shown in the GUI. In general, this algorithm contributes to the overall enhancement of the identity verification pipeline in that the presentation attacks based on the analysis of the texture patterns, compression artifacts, and the nature of image capture are detected. The originality of the proposed structure lies in the fact that three complementary verification dimensions have been integrated into one real-time structure. It utilizes first deep structural similarity learning via a Siamese CNN to conduct document verification in embedding space in a template manner. Second, it proposes country-sensitive adaptive fraud thresholding, which allows decision-making to be carried out with a calibration that is consistent with national document standards and variability. Third, it uses an EfficientNet-based presentation attack detector to differentiate between original captures and scanned, or digitally re-created document based on forensic texture analysis. The framework overcomes the weaknesses of single-modality systems by using the structural verification, adaptive decision-level classification, and capture-modality analysis in a single deployment pipeline. This combined design is stronger, more scalable, and, for remote onboarding and high-security digital identity settings, it is more focused.

The three modules are intended to be self-contained checks. The similarity score is reported in Module 1, fraud classifications are conducted in Module 2 with adaptive thresholds and presentation attack detection is done in Module 3. The current implementation does not automatically fuse decisions from each of the modules but instead provides all the outputs, one per module, to the operator. Each verification function can be used independently of the others, as per the application. The results will be analyzed for confidence-based decision fusion, which can be applied to integrate the outputs into a single authentication score in the future.

#### 4. Experimental Setup

In this section, the three proposed modules 1 and 2 will be described, as they have a similar architecture and training setups, hence are collectively described. The separate presentation of Module 3, which aims at a different detection problem, is given. The data, preprocessing, model configuration, software platform, and metrics of evaluation are illustrated in this section.

##### 4.1. Experimental Setup for Module 1 and Module 2

Experimental evaluation was made with the help of the MIDV-2020 dataset [30]. This data set has various categories of identity documents that were issued in 10 countries.

Training set: 320 genuine identity documents and 320 forged identity documents per country. Test sample: 80 genuine and 80 forged documents from a country. Generation of forgeries: The generation of forged samples was carried out based on a copy-move attack on the  $64 \times 64$ -pixels regions to emulate the forgeries.

All images of the documents were downsized  $300 \times 300$  and scaled down after which they were passed through the modules. Siamese network training was done through the construction of pairs of images that comprised of genuine-genuine and genuine-forged combinations.

The model format of modules 1 and 2 includes a Siamese neural network, which is trained on a contrastive loss. The network was trained at a batch size of 4, and 2000 epochs were used to update the weights of the network using the Adam optimizer. Each model had a country-specific learning rate as

presented in Table 2 to allow optimal convergence and performance of the identity documents of every specific country.

**Table 2.** Training Parameters and Final Loss Values per Country

| Country | True IDs | Forged IDs | LR                 | Epochs | Batch | Loss                 |
|---------|----------|------------|--------------------|--------|-------|----------------------|
| Alb     | 320      | 320        | $5 \times 10^{-5}$ | 2000   | 4     | $1.6 \times 10^{-3}$ |
| Aze     | 320      | 320        | $5 \times 10^{-4}$ | 2000   | 4     | $3.8 \times 10^{-4}$ |
| Esp     | 320      | 320        | $5 \times 10^{-5}$ | 2000   | 4     | $7.5 \times 10^{-5}$ |
| Est     | 320      | 320        | $5 \times 10^{-4}$ | 2000   | 4     | $3.0 \times 10^{-4}$ |
| Fin     | 320      | 320        | $5 \times 10^{-4}$ | 2000   | 4     | $4.9 \times 10^{-4}$ |
| Grc     | 320      | 320        | $5 \times 10^{-4}$ | 2000   | 4     | $4.2 \times 10^{-4}$ |
| Iva     | 320      | 320        | $5 \times 10^{-4}$ | 2000   | 4     | $4.7 \times 10^{-4}$ |
| Rus     | 320      | 320        | $5 \times 10^{-4}$ | 2000   | 4     | $2.7 \times 10^{-3}$ |
| Srb     | 320      | 320        | $5 \times 10^{-4}$ | 2000   | 4     | $1.2 \times 10^{-3}$ |
| Svk     | 320      | 320        | $5 \times 10^{-4}$ | 2000   | 4     | $2.1 \times 10^{-4}$ |

The training setup and the final loss of each country-specific model are given in Table 2. The dataset was balanced so that each country had 320 genuine samples and 320 forged identity samples. The data was trained using 2000 epochs with a batch size of 4. Two learning rates were used ( $5 \times 10^{-5}$  and  $5 \times 10^{-4}$ ), which were chosen empirically to maintain the stability of the convergence. The findings show that the convergence between all models is similar, and the final values of losses are small (between  $7.5 \times 10^{-5}$  (Spain) and  $2.7 \times 10^{-3}$  (Russia)).

#### 4.1.1. Scalability of the Country-Specific Models

The current approach is to have a separate model for each country, but this was chosen to assess the proposed approach under controlled experimental conditions and to include the large structural differences between the identity document templates in the MIDV-2020 dataset. By having similar thresholds tailored to each country, and a more robust ability to handle the variety of security graphics in document layouts, country-specific training can be done.

However, from a deployment point of view, keeping a separate model for each country or document type would become more and more difficult as the number of supported documents increased. A possible more scalable implementation would be a single shared backbone network trained on documents from several countries, with light-weight country-specific calibration parameters or country-specific adaptive decision thresholds. In this format, the feature extraction network would be common for all documents, while just the decision layer or threshold values would be customized for each country, significantly reducing storage space, maintenance, and retraining expenses.

This evolution of course is well served by the proposed framework, as thresholds for each country are already distinct from the feature extraction process. In future work, we shall explore unified multi-country learning, domain adaptation and continual learning strategies to allow the framework to learn new document types without retraining an independent model for each country.

In this study, the 70/30 train-test split was applied to each country separately as the goal was to assess the country specific models with the adaptive threshold for heterogeneous layouts of documents. However, the reported results are not an assessment of generalization across the country. Leave-one-country-out evaluation (training models from several countries and testing them in unseen countries) and domain adaptation (modifying the model to improve its performance when the data domain is shifted) are future directions for this work that will be explored.

The verification system is tested with three major measures:

Equation (1) represents the True Acceptance Rate (TAR): This is the ratio of the number of accepted pairs of identity documents that are similar.

$$TAR = \frac{\text{True Positives (TP)}}{\text{True Positives (TP)} + \text{False Negatives (FN)}} \quad (1)$$

Equation (2) represents the False Rejection Rate (FRR): Measures the proportion of genuine identity document pairs incorrectly rejected.

$$FRR = \frac{\text{False Negatives (FN)}}{\text{True Positives (TP)} + \text{False Negatives (FN)}} \quad (2)$$

Equation (3) represents the False Acceptance Rate (FAR): Measures the proportion of forged identity document pairs incorrectly accepted as genuine.

$$FAR = \frac{\text{False Positives (FP)}}{\text{False Positives (FP)} + \text{True Negatives (TN)}} \quad (3)$$

These metrics are calculated using a binary decision threshold that is used on the similarity scores produced by the Siamese neural network. Each country has a threshold that is used to strike a balance between FAR and TAR.

#### 4.2. Experimental Setup for Module 3

Module 3 was also tested on the MIDV-2020 dataset, and the results were evaluated in terms of detecting scanned counterparts of authentic identity documents. In every target country, the data were composed of 100 authentic documents and 100 scans, and it was broken down into 70% training and 30% testing. To achieve a comprehensive and systematic assessment, training pairs were designed to have a total of 14,560 combinations in each country, which comprise 4,830 real-real pairs, 4,900 real-scan pairs, and 4,830 scan-scan pairs. The test set contained 2,640 pairs in each country with equal distribution as 870 real-real, 900 real-scan, and 870 scan-scan combinations. The model was trained on the Adam optimizer, with an initial learning rate of 10 (decayed after 10 epochs) out of 100 total epochs, in total to ensure that progress is stable. It is a systematic pairing approach that ensures that the model is extensively tested to determine its potential to differentiate between original documents and high-quality scans regardless of fine-grained security features, and that the model will perform well in all types of documents and countries.

The balanced pairing approach was designed to avoid the classifier being biased towards the majority class and to provide robust training for the features of the original and scanned documents. This experimental setup gives a good baseline for comparing performance of models in different countries, but real-world identity verification systems may have a far higher percentage of genuine documents relative to fraudulent or scanned documents. The proposed framework will therefore be evaluated in future under realistic class-imbalanced settings and the false-positive rate will be analyzed.

### 5. Experimental Results

The results of the performance of the verification framework using module 1 and module 2 are reported in this section. The training of both models was done under the same training protocol and dataset configuration as described in the preceding parts of the experiment. The performance is estimated using distance scores that are obtained by the Siamese network and relevant measures such as TAR, FRR, and FAR are calculated. The findings have been reported in each country, and the level of differentiation of the two models on the authentic and counterfeit pairs of identity documents at different levels of threshold is indicated.

#### 5.1. Experimental Results for Module 1 and Module 2

To test the efficiency of the Siamese network, the similarity scores of pairs of images are computed. Table 3 presents the information about the results of testing, which includes the TAR, FRR, FAR, and the similarity distances of each country.

**Table 3.** Testing Results: Similarity Distances and Evaluation Metrics.

| Country | Threshold   | TAR  | FRR   | FAR   | Min Dist. | Max Dist. | Min Dist. | Max Dist. |
|---------|-------------|------|-------|-------|-----------|-----------|-----------|-----------|
| Alb     | [0.9–1.10]  | 0.91 | 0.085 | 0.14  | 0.00      | 1.83      | 0.05      | 2.70      |
| Aze     | [0.75–1.25] | 0.95 | 0.050 | 0.073 | 0.00      | 1.71      | 0.023     | 2.83      |
| Esp     | [0.5–1.00]  | 0.99 | 0.010 | 0.06  | 0.00      | 1.92      | 0.06      | 2.10      |
| Est     | [0.75–1.00] | 0.89 | 0.100 | 0.00  | 0.00      | 1.30      | 0.63      | 2.66      |
| Fin     | [1.0–1.10]  | 0.99 | 0.013 | 0.00  | 0.00      | 0.97      | 1.16      | 2.63      |
| Grc     | [1.0–1.10]  | 0.95 | 0.045 | 0.027 | 0.00      | 1.49      | 0.29      | 2.67      |
| Iva     | [1.0–1.10]  | 0.96 | 0.040 | 0.020 | 0.00      | 2.11      | 0.10      | 2.66      |
| Rus     | [1.0–1.25]  | 0.91 | 0.090 | 0.040 | 0.00      | 1.21      | 0.14      | 2.53      |
| Srb     | [0.75–0.95] | 1.00 | 0.000 | 0.030 | 0.00      | 0.82      | 0.14      | 2.60      |

|     |             |      |       |       |      |      |      |      |
|-----|-------------|------|-------|-------|------|------|------|------|
| Svk | [0.80–1.10] | 1.00 | 0.000 | 0.030 | 0.00 | 0.75 | 1.15 | 2.64 |
|-----|-------------|------|-------|-------|------|------|------|------|

Table 3 presents the similarity-based verification performance across different countries using adaptive threshold ranges. Overall, the proposed model achieves high TAR, ranging from 0.89 to 1.00, with several countries (e.g., Srb and Svk) reaching perfect acceptance (TAR = 1.00). The FRR remains low (0.000–0.100), while the FAR is generally limited (0.00–0.14), indicating effective discrimination between genuine and forged ID pairs.

The model had high TARs in most of the countries, as shown in Table 3, especially in Spain, Serbia, and Slovakia, where the TAR was 99–100. Low FARs indicate, in general, a robustness against false matches. The threshold ranges were adjusted to give an equal balance between TAR and FAR. The verification system that was established in Siamese showed high capability in the process of differentiating the authentic and counterfeited documents.

After training Siamese network, the thresholds were determined individually for each country, depending on the validation data. A similar distance distribution was used to assess candidate thresholds and an operating threshold was chosen based on the optimization criterion given in Section 5.1.1. The reported intervals represent operating ranges of performance where the verification performance varied by a small degree and showed that the decision boundaries varied with different document types.

To further test the performance of the proposed detection framework, a real-time experiments on two different datasets that have been used is performed: the first dataset is the real identity document images, which were obtained in MIDV-500 and Wikimedia Commons. The second data is the artificially forged documents that have been generated by a copy-move manipulation paradigm with  $64 \times 64$ -pixel patches. The capacity to assess ten countries was broad in terms of the ID card designs and security features incorporated. Figure 2 shows the way the models can manage various document types and forgery cases. This real-time validation provides the resistance and flexibility of the proposed verification system in the real conditions.



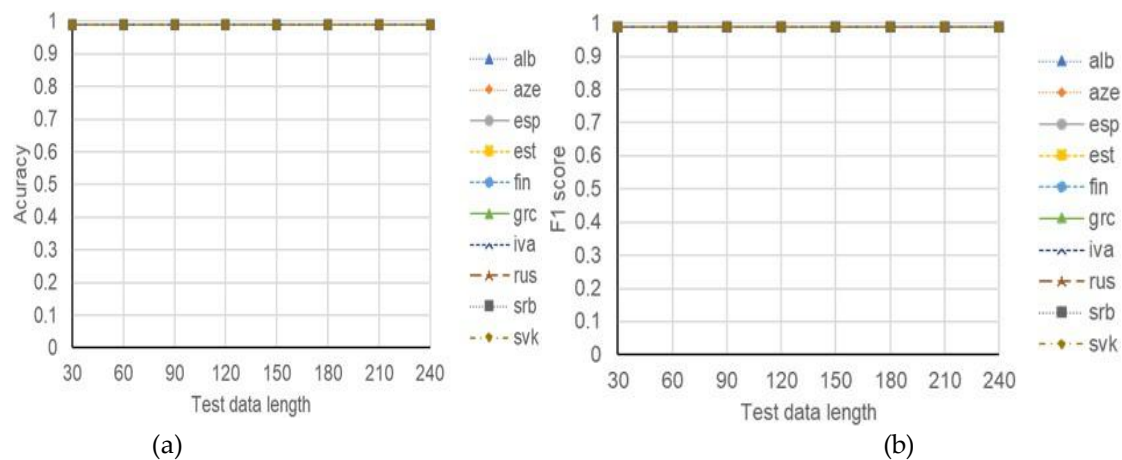
**Figure 2.** Real-time experimental results across ten countries using diverse identity cards from the MIDV2020 dataset, demonstrating the model's performance in verifying genuine and forged documents.

### 5.1.1. Threshold Calibration

The validation subset of training data was used to determine the decision threshold for each country. Similarity distance between genuine-falsified pairs were calculated, after which candidate thresholds were assessed with reference to the observed similarity distance distribution. The threshold chosen is the operating point that maximizes the True Acceptance Rate (TAR) and the False Acceptance Rate (FAR) while minimizing false classifications. Optimal operating point varies by country due to variations in document layouts and similarity distributions. As a result, Table 3 shows the thresholds that kept this balance stable in each country.

### 5.2. Experimental Results for Model 3

The outcomes of the evaluation prove that module 3 demonstrates a high level of performance in all the countries under evaluation, as well as in test data of different lengths. Both the accuracy and F1- scores are nearly 1.0 for all test sizes (Figures 3a and 3b), which indicates good generalization and strength.



**Figure 3.** Model 3 performance in terms of evaluation on the MIDV2020 dataset with various test data lengths. (a) The trend of accuracy per country has a constant, near perfection performance. (b) The tendency of the F1-score, which shows the stability of classification of all test sets.

Real-time verification of the original/scanned document detecting module shows that the model has the capability to identify with the correct identity documents, as well as the scanned copies. The system is very accurate, as illustrated by Figure 4, even under different lighting conditions and the quality of an image. This strength is vital in real-life applications where the input images can be taken using various devices and under uncontrolled conditions. Results prove the appropriateness of the model to the maintenance of the fine texture and compression artifact information that differentiates genuine documents with scans to facilitate the detection of fraud in the workplace.

| Country | Real Samples (MIDV2020) | Detect Real? | Scan Samples (MIDV2020) | Detect Scan? |
|---------|-------------------------|--------------|-------------------------|--------------|
| Alb     |                         | ✓            |                         | ✓            |
|         |                         | ✓            |                         | ✓            |
| Aze     |                         | ✓            |                         | ✓            |
|         |                         | ✓            |                         | ✓            |
| Esp     |                         | ✓            |                         | ✓            |
|         |                         | ✓            |                         | ✓            |
| Est     |                         | ✓            |                         | ✓            |
|         |                         | ✓            |                         | ✓            |

**Figure 4.** Real-time validation results for original vs. scanned identity document detection.

### 5.3. System Performance and Efficiency Analysis

To evaluate the practicality of deploying real-world applications, especially for edge devices and mobile/IoT systems, the main performance metrics are measured, such as latency, throughput, and model efficiency, across the three modules. Measurements were conducted on an Intel i7 CPU (simulating edge scenarios) and Tesla V100 GPU (server deployment), using batch size=4 and PyTorch with `torch.no_grad()`

for optimized inference. Table 4 presents the latency results (*ms* per document), where all timings represent averages over 100 runs with standard deviations.

**Table 4.** Latency results (*ms* per document).

| Module                                                             | CPU<br>(Intel i7) | GPU<br>(Tesla V100) | CPU Throughput<br>(docs/sec) | GPU Throughput<br>(docs/sec) |
|--------------------------------------------------------------------|-------------------|---------------------|------------------------------|------------------------------|
| <b>Module 1: Similarity-<br/>Based Structural<br/>Verification</b> | 11.30 ms          | 2.17 ms             | 88.6                         | 459.7                        |
| <b>Module 2: Country-<br/>Specific Fraud<br/>Detection</b>         | 11.77 ms          | 2.30 ms             | 85.0                         | 434.7                        |
| Module 3: Original vs.<br>Scanned Document<br>Classification       | 45.60 ms          | 5.90 ms             | 21.9                         | 169.4                        |
| Full Pipeline                                                      | 69.62 ms          | 10.52 ms            | 14.4                         | 95.0                         |

The throughput ( $throughput = 1000 / reported\_latency$ ) using CPU/GPU is also calculated. The throughput using CPU is equal to 14.4 docs/sec (pipeline), and using the GPU is equal to 94.8 docs/sec (~24 fps, suitable for real-time GUI operation), where:

$$Frame\ per\ second = \frac{doc/sec}{batch\ size} = \frac{95.0}{4} \approx 24\ fps$$

Additionally, the model efficiency metrics are presented. For the Siamese CNN, the model efficiency is equal to 2.1M parameters, 8.6 MB (lightweight, edge-deployable). And for EfficientNet-B7, the model efficiency is equal to 66.4M parameters, 268 MB (pruning-ready, fine-tuned)

The Siamese CNN's compact architecture (3 conv blocks + 5D embeddings from  $300 \times 300$  grayscale inputs) enables sub-50ms inference on CPU, while EfficientNet-B7 handles texture analysis efficiently despite higher latency. Timings showed GUI integration, confirming less than 300ms end-to-end on normal hardware. To achieve mobile optimization, it will be optimized next with INT8 quantization (should be 4x faster) and 30-50% pruning with less than 1% accuracy loss.

## 6. Comparative Study and Discussion

Table 5 provides a comparison of the capabilities of the proposed framework with those of representative ID document verification methods. CheckSim [15] is one of the most pertinent studies, which deals with reference-based document similarity verification using Siamese and Triplet neural networks; CheckScan [18] is another study that deals with scanned document attacks, using a frequency-domain analysis. DLD layout descriptor [19] suggested a structural verification method using document layout analysis. While these approaches address some aspects of document verification, they each focus on a specific problem within document verification.

**Table 5.** Functional comparison with representative methods of ID verification.

| Method                        | Structural<br>Verification | Fraud<br>Detection | Scan<br>Detection | Adaptive<br>Thresholds | Multi-<br>Module | Real-<br>time |
|-------------------------------|----------------------------|--------------------|-------------------|------------------------|------------------|---------------|
| CheckSim [15]                 | ✓                          | ✗                  | ✗                 | ✗                      | ✗                | ✓             |
| CheckScan [18]                | ✗                          | ✗                  | ✓                 | ✗                      | ✗                | ✓             |
| DLD layout<br>descriptor [19] | ✓                          | ✗                  | ✓                 | ✗                      | ✗                | ✗             |
| <b>Proposed work</b>          | ✓                          | ✓                  | ✓                 | ✓                      | ✓                | ✓             |

Compared to existing approaches for solving individual verification tasks, the proposed framework combines structural verification, adaptive fraud detection and presentation attack detection in a single real-time pipeline. Table 6 also summarizes the datasets and evaluation protocols considered in the representative approaches to further compare with previous work.

**Table 6.** A comparative study between the proposed work and the other related work in literature.

| Method               | Main Task                                                                   | Dataset   | Reported Performance                                                               |
|----------------------|-----------------------------------------------------------------------------|-----------|------------------------------------------------------------------------------------|
| CheckScan [18]       | Scan detection                                                              | MIDV-2020 | 93–100% country-wise detection accuracy depending on document type and hash length |
| <b>Proposed work</b> | Structural verification, fraud detection, and presentation attack detection | MIDV-2020 | <b>97.8% Accuracy, 99.0% TAR, 3.4% FAR, F1-score = 0.98</b>                        |

As illustrated in table 6, the direct comparison of the proposed framework with previous study is not statistically valid because different verification objectives were considered, and different experimental protocols and evaluation metrics were used. For instance, CheckScan only focuses on detection of scan attacks and provides country wise detection accuracy report. The reported values should, therefore, not be interpreted as direct comparisons.

However, the comparison shows the proposed framework can solve a more comprehensive verification problem, which includes multiple verification modules into one verification pipeline and keeps competitive performance on the publicly available MIDV-2020 benchmark. Such integration is novel as compared to the existing methods that have been focused on performing only one verification task, by performing the three of them simultaneously, the proposed approach can provide structural verification, fraud detection and presentation attack detection.

## 7. User Interface Design (Prototype)

The prototype GUI is also structured into three main modules, which facilitate the identity document verification process. The home page offers direct access to the basic functionalities such as source checking, detection of fraud, detection of original versus scanned document, and tracking. All modules are developed to allow an easy interaction with the users, and document selection is made possible through drop-down menus, and visual comparison is displayed side by side.

The source verification module calculates a Guilloche-based similarity score between a control document and a reference document, with a lower value suggesting that the two are more visually similar. The fraud detection module identifies the document as a fake or authentic document using the learned patterns of the verification models. The original scanned detection module, which is the implementation of module 3, assesses whether the document is original or scanned depending on a structural and visual similarity score.

### 7.1. Main Page

The following figure shows the main page of the verification tool for IDs. This page provides navigation buttons for the different functionalities: Figure 5 presents the main page of the verification tool. This page provides quick access to the core functionalities through the following navigation buttons:

- Source verification: a module for calculating the visual similarity between a control ID and a reference.
- Fraud detection: module for detecting forged or manipulated identity documents.
- Original/Scan detection: determines whether the input document is an original or a scanned copy.
- Tracking: used to track the ID in video in the wild.
- Reset: clears all selections and resets the workspace.

### 7.2. Source Verification Module

This module is used to assess the visual similarity between a reference identity document and a control document. It provides the foundation for further fraud analysis. The user can select documents via dropdown menus and visualize them side by side. This is because a Guilloche-based similarity score is determined, and a score near zero means that there is a high level of similarity. Figure 6 represents the source verification module that compares a reference ID and a control ID.

### 7.3. Fraud Detection Model

The module of fraud detecting will detect forged identity documents. The system, as illustrated in Figure 7, uses two identity documents as input and categorizes one as fake and the other one as real, relying upon learned patterns in the verification model.

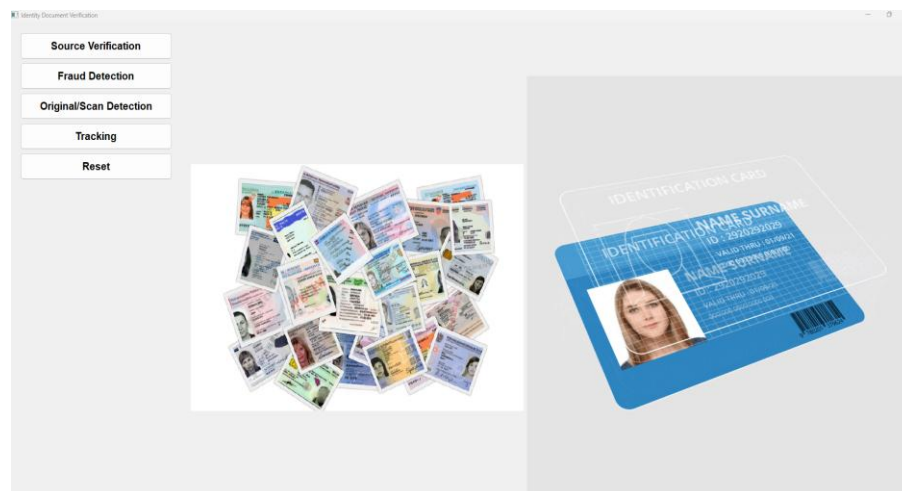


Figure 5. Main page of the identity verification tool.

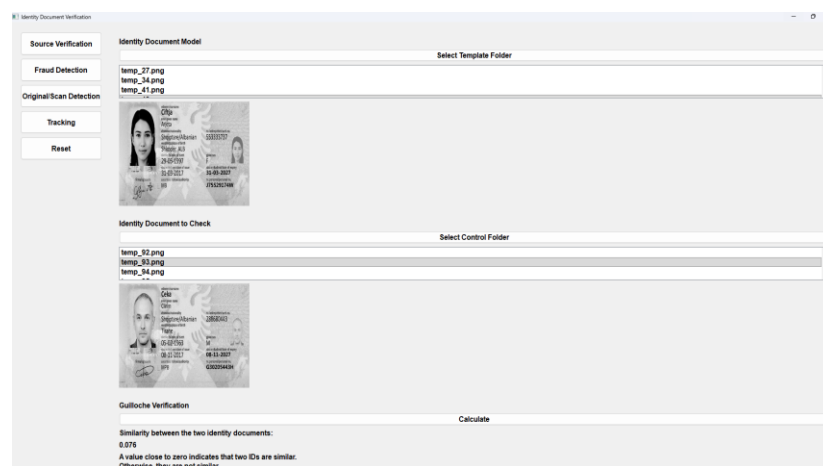
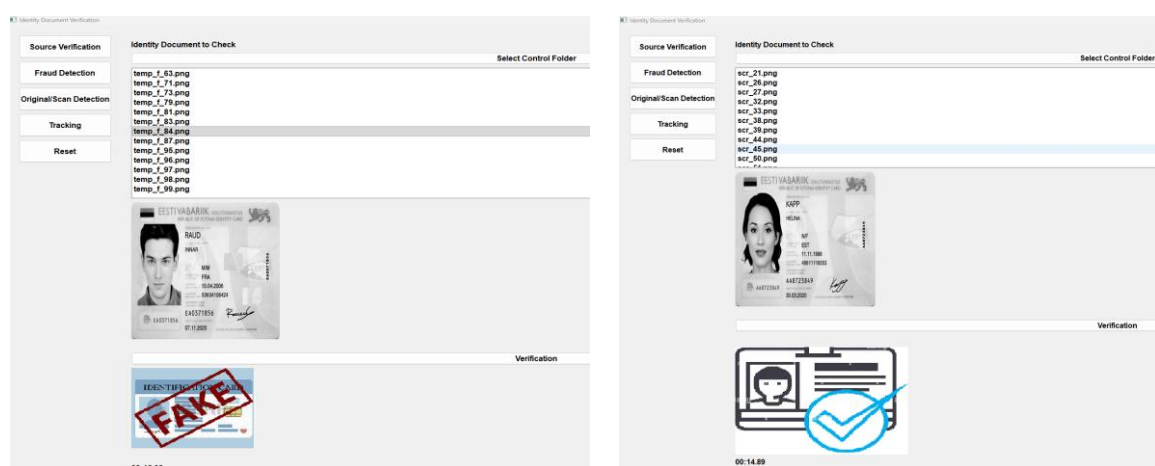


Figure 6. Source verification module comparing a reference ID and a control ID.



(a) Fake ID classification result.

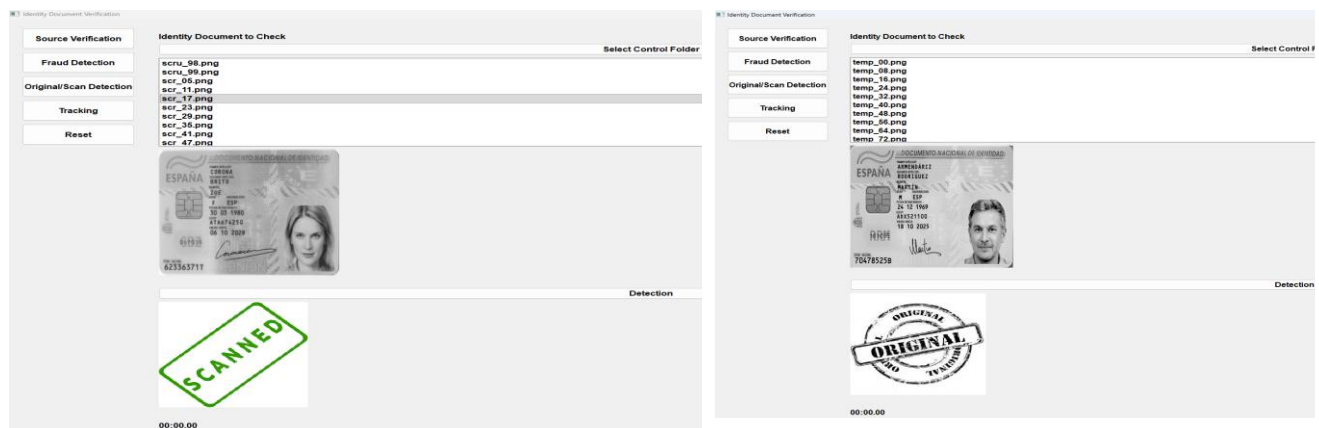
(b) Real ID classification result.

Figure 7. Interface of the fraud detection module.

#### 7.4. Original vs. Scanned Detection Module

The current module deploys the Model 3 that seeks to identify whether a document is a scanned or an original one. As shown in Figure 8, the system allows the user to choose a reference document and a control document. In comparison with it, it produces a similar score. A lesser score suggests that both

documents are structurally and visually alike (i.e., like both originals), whereas a higher score is an indication of a scanned or altered document.



(a) Detection of a scanned identity document.

(b) Detection of an original identity document.

**Figure 8.** Original vs. Scanned ID detection module.

To complement the prototype description, the interface was evaluated quantitatively using a subset of the MIDV-2020 dataset. Table 7 summarizes the performance metrics for each module across 500 document pairs.

**Table 7.** Quantitative Evaluation of the Prototype Modules on 500 document pairs.

| Module                                   | Metric                      | Value |
|------------------------------------------|-----------------------------|-------|
| Similarity-Based Structural Verification | Similarity Accuracy (%)     | 94.2% |
|                                          | Precision (%)               | 96.8% |
|                                          | Recall (%)                  | 94.7% |
| Country-Specific Fraud Detection         | True Acceptance Rate (%)    | 95.6% |
|                                          | False Acceptance Rate (%)   | 3.4%  |
|                                          | Classification Accuracy (%) | 96.1% |
|                                          |                             |       |
| Original/Scan Detection                  | Precision (%)               | 95.9% |
|                                          | Recall (%)                  | 96.4% |

Table 7 shows the quantitative performance of the proposed prototype in the three fundamental verification modules of the prototype. The similarity-based structural verification module had a high similarity accuracy of 94.2% because it has been shown to have identified structural similarity between identity documents and reference templates. A good result was obtained with the country-specific fraud detection module, having a high precision of 96.8%, a recall of 94.7%, and a true acceptance of 95.6% with a low false acceptance rate of 3.4%, which is a great indicator of discriminating accuracy between the real documents and the fraudulent documents. Also, original versus scanned document detection module had a classification accuracy of 96.1%, precision, and recall of 95.9% and 96.4%, respectively. Overall, the findings attest to the adequacy and stability of the suggested multi-module model in the correct authentication of the documents and detection of fraud at the various verification levels.

## 8. Conclusion and Future Work

This paper introduced a multimodal identity document verification framework, which increases the security of remote onboarding by combining the structural verification, fraud detection and presentation attack detection into a single framework. The proposed system includes Siamese CNN for visual similarity analysis, country-specific adaptive fraud thresholds and the EfficientNet-B7 classifier to classify original documents and scanned copies. This holistic design will overcome some of the major drawbacks of the stand-alone designs and will enable for better Digital identity assurance in environments with requirements to control. The experimental results with the MIDV-2020 dataset showed excellent results for all modules. The similarity-based verification and fraud detection modules obtained the TAR in the range

of 0.89 to 1.00 and FAR in the range of 0.00 to 0.14. The accuracy of the prototype structural verification module was 94.2%, the precision of the prototype fraud detection module was 96.8%, the prototype recall was 94.7%, the prototype TAR was 95.6%, and the prototype FAR was 3.4% in 500 document pairs test. The accuracy of the original-versus-scanned classification module was 96.1%, the precision was 95.9%, and the recall was 96.4% in the 500 document pairs test. The entire verification pipeline was also tested and shown to be real-time: the average verification time was 69.62 ms for CPU and 10.52 ms for GPU. The framework is expected to be further developed and validated with larger and more heterogeneous real-world datasets, expanded to accommodate more document types and countries, and further enhanced to withstand more difficult imaging conditions and new AI created documents. Future studies will also consider the framework in realistic class-imbalanced operational configurations, and explore the possibilities of generalization across countries using a unified multi-country training and leave-one-country-out evaluation approach.

**Conflicts of Interest:** “The authors declare no conflict of interest

## References

1. Muñoz-Haro, J., Tolosana, R., Fierrez, J., Vera-Rodriguez, R., & Morales, A. (2026). Privacy-aware detection of fake identity documents: methodology, benchmark, and improved algorithms (FakeIDet2). *Information Fusion*, 128, 103969. <https://doi.org/10.1016/j.inffus.2025.103969>
2. Voerman, J., Al-Ghadi, M., Sidere, N., Coustaty, M., & Lessard, O. (2025). Optimizing identity documents classification in online systems: a comparative analysis. *International Journal on Document Analysis and Recognition (IJДАР)*. <https://doi.org/10.1007/s10032-025-00555-5>
3. Xie, L., Wang, Y., Guan, H., Nag, S., Goel, R., Swamy, N., Yang, Y., Xiao, C., Prisby, J., Maciejewski, R., & Zou, J. (2024, December). IDNet: A Novel Identity Document Dataset via Few-Shot and Quality-Driven Synthetic Data Generation. In *2024 IEEE International Conference on Big Data (BigData)* (pp. 2244–2253). <https://doi.org/10.1109/BigData62323.2024.10825017>
4. Abdallah, A., Maarof, M. A., & Zainal, A. (2016). Fraud detection system: A survey. *Journal of Network and Computer Applications*, 68, 90–113. <https://doi.org/10.1016/j.jnca.2016.04.007>
5. Lugon-Moulin, S., Weyermann, C., & Baechler, S. (2022). An efficient method to detect series of fraudulent identity documents based on digitized forensic data. *Science & Justice*, 62, 610–620. <https://doi.org/10.1016/j.scijus.2022.09.003>
6. Al-Ghadi, M., Ming, Z., Gomez-Krämer, P., Burie, J.-C., Coustaty, M., & Sidere, N. (2023). Guilloche detection for ID authentication: A dataset and baselines. In *Proceedings of the International Workshop on MultiMedia Signal Processing (MMSP)*, 1–6. IEEE. <https://doi.org/10.1109/MMSP59012.2023.10337681>
7. Al-Ghadi, M., Voerman, J., Coustaty, M., Lessard, O., & Sidere, N. (2025). IDTrust: Deep identity document quality detection with bandpass filtering. In *Proceedings of the Winter Conference on Applications of Computer Vision (WACV) Workshops*, 716–723. IEEE. <https://doi.org/10.1109/WACVW65960.2025.00081>
8. Kada, O., Kurtz, C., Kieu, C., & Vincent, N. (2022). Hologram detection for identity document authentication. In El Yacoubi, M., Granger, E., Yuen, P. C., Pal, U., & Vincent, N. (eds.) *Pattern Recognition and Artificial Intelligence*, 13363, 346–357. Springer, Cham. [https://doi.org/10.1007/978-3-031-09037-0\\_29](https://doi.org/10.1007/978-3-031-09037-0_29)
9. Chapel, M.-N., Al-Ghadi, M., & Burie, J.-C. (2023). Authentication of holograms with mixed patterns by direct LBP comparison. In *Proceedings of the International Workshop on MultiMedia Signal Processing (MMSP)*, 7–12. IEEE. <https://doi.org/10.1109/MMSP59012.2023.10337669>
10. Pouliquen, G., Chiron, G., Chazalon, J., Géraud, T., & Awal, A. M. (2024). Weakly supervised training for hologram verification in identity documents. In E. H. Barney Smith, M. Liwicki, & L. Peng (Eds.), *Document Analysis and Recognition - ICDAR 2024*, 14804, 17–33. Springer, Cham. [https://doi.org/10.1007/978-3-031-70533-5\\_2](https://doi.org/10.1007/978-3-031-70533-5_2)
11. Pouliquen, G., Chazalon, J., Chiron, G., Géraud, T., & Awal, A. M. (2026). Verification of dynamic holographic behavior in identity documents. In X. C. Yin, D. Karatzas, & D. Lopresti (Eds.), *Document Analysis and Recognition ICDAR 2025*, 16025, 323–339. Springer, Cham. [https://doi.org/10.1007/978-3-032-04624-6\\_19](https://doi.org/10.1007/978-3-032-04624-6_19)
12. Xu, J., Jia, D., Lin, Z., Zhou, T., Wu, J., & Tang, L. (2024). PBNet: Combining Transformer and CNN in passport background texture printing image classification. *Electronics*, 13(21), 4160. <https://doi.org/10.3390/electronics13214160>
13. Zhu, Y., Wang, H., Qian, Z., Li, S., Zhang, X., & Liu, J. (2025). Towards generalized physical occlusion detection on documents. In *MM '25: Proceedings of the 33rd ACM International Conference on Multimedia* (pp. 7596–7605). ACM. <https://doi.org/10.1145/3746027.3754975>
14. Wang, H., Li, S., Cao, S., Yang, R., Zeng, J., Qian, Z., & Zhang, X. (2023, October). On physically occluded fake identity document detection. In *MM '23: Proceedings of the 31st ACM International Conference on Multimedia* (pp. 1556–1564). ACM. <https://doi.org/10.1145/3581783.361207>
15. Ghanmi, N., Nabli, C., & Awal, A. M. (2021). CheckSim: A reference-based identity document verification by image similarity measure. In Barney Smith, E. H., & Pal, U. (eds.) *Document Analysis and Recognition – ICDAR 2021 Workshops*, 12916, 422–436. Springer, Cham. [https://doi.org/10.1007/978-3-030-86198-8\\_30](https://doi.org/10.1007/978-3-030-86198-8_30)
16. Dey, S., Dutta, A., Toledo, J., Ghosh, S., Lladós, J., & Pal, U. (2017). SigNet: Convolutional Siamese Network for writer independent offline signature verification. *arXiv:1707.02131*. <https://doi.org/10.48550/arXiv.1707.02131>
17. Centeno, A. B., Terrades, O. R., Canet, J. L., & Morales, C. C. (2019). Recurrent comparator with attention models to detect counterfeit documents. In *Proceedings of the International Conference on Document Analysis and Recognition (ICDAR)*, 1332–1337. IEEE Computer Society, Sydney, Australia. <https://doi.org/10.1109/ICDAR.2019.00215>

18. Al-Ghadi, M., Gomez-Krämer, P., & Burie, J.-C. (2022). CheckScan: a reference hashing for identity document quality detection. In Osten, W., Nikolaev, D., & Zhou, J. (eds.) *14th International Conference on Machine Vision (ICMV 2021)*, 120840. SPIE, Rome, Italy. <https://doi.org/10.1117/12.2623887>
19. Gomez-Krämer, P., Rouis, K., Diallo, A. O., & Coustaty, M. (2024). Printed and scanned document authentication using robust layout descriptor matching. *Multimedia Tools and Applications*, 83(16), 47477–47502. <https://doi.org/10.1007/s11042-023-17021-1>
20. Tapia Farias, J., & Busch, C. (2025). Can foundation models generalise the presentation attack detection capabilities on ID cards? arXiv preprint arXiv:2506.05263. <https://doi.org/10.48550/arXiv.2506.05263>
21. Castelblanco, A., Solano, J., Lopez, C., Rivera, E., Tengana, L., & Ochoa, M. (2020). Machine learning techniques for identity document verification in uncontrolled environments: A case study. In Figueroa Mora, K., Anzures Marín, J., Cerda, J., Carrasco-Ochoa, J., Martínez-Trinidad, J., & Olvera-López, J. (eds.) *Pattern Recognition. MCPR 2020*, 12088, 271–281. Springer, Cham. [https://doi.org/10.1007/978-3-030-49076-8\\_26](https://doi.org/10.1007/978-3-030-49076-8_26)
22. Zhang, C. J., Gill, A., Liu, B., & Anwar, M. J. (2025). AI-based identity fraud detection: A systematic review. *ArXiv*, abs/2501.09239. <https://doi.org/10.48550/arXiv.2501.09239>
23. Sukhija, R., Kumar, M., & Jindal, M. K. (2025). Document forgery detection: a comprehensive review. *International Journal of Data Science and Analytics*, 20(5), 4385–4407. <https://doi.org/10.1007/s41060-025-00723-0>
24. Chinapas, A., Polpinit, P., & Saikaew, K. (2019). Personal verification system using the ID card and face photo for cross-age face. In *ICSEC 2019 - 23rd International Computer Science and Engineering Conference*, 325–330. Institute of Electrical and Electronics Engineers Inc., Phuket, Thailand. <https://doi.org/10.1109/ICSEC47112.2019.8974845>
25. Wu, X., Xu, J., Wang, J., Li, Y., Li, W., & Guo, Y. (2019). Identity authentication on mobile devices using face verification and ID image recognition. *Procedia Computer Science*, 162, 932–939. <https://doi.org/10.1016/j.procs.2019.12.070>
26. Bertrand, R., Gomez-Krämer, P., Terrades, O. R., Franco, P., & Ogier, J.-M. (2013). A system based on intrinsic features for fraudulent document detection. In *12th International Conference on Document Analysis and Recognition*, 106–110. <https://doi.org/10.1109/ICDAR.2013.29>
27. Martínez-Tornés, B., Boros, E., Doucet, A., Gomez-Krämer, P., & Ogier, J.-M. (2023). Detecting forged receipts with domain-specific ontology-based entities & relations. In Fink, G. A., Jain, R., Kise, K., & Zanibbi, R. (eds.) *Document Analysis and Recognition - ICDAR 2023*, 14189, 184–199. Springer, Cham. [https://doi.org/10.1007/978-3-031-41682-8\\_12](https://doi.org/10.1007/978-3-031-41682-8_12)
28. Ghanmi, N., & Awal, A. M. (2018). A new descriptor for pattern matching: application to identity document verification. In *Proceedings - 13th IAPR International Workshop on Document Analysis Systems (DAS 2018)*, 375–380. Institute of Electrical and Electronics Engineers Inc., Vienna, Austria. <https://doi.org/10.1109/DAS.2018.74>
29. Talbot-Wright, B., Baechler, S., Morelato, M., Ribaux, O., & Roux, C. (2016). Image processing of false identity documents for forensic intelligence. *Forensic Science International*, 263, 67–73. <https://doi.org/10.1016/j.forsciint.2016.03.054>
30. Bulatov, K., Emelianova, E., Tropin, D., Skoryukina, N., Chernyshova, Y., Sheshkus, A., Usilin, S., Ming, Z., Burie, J.-C., Luqman, M. M., & Arlazarov, V. (2022). MIDV-2020: a comprehensive benchmark dataset for identity document analysis. *Computer Optics*, 46, 252–270. <https://doi.org/10.18287/2412-6179-CO-1006>
31. ALDabbas, A., Baniata, L. H., AlSaaidah, B. A., Mustafa, Z., Alali, M., & Rateb, R. (2025). Artificial intelligence-driven method for the discovery and prevention of distributed denial of service attacks. *Int J Artif Intell ISSN*, 2252(8938), 8938. <https://doi.org/10.11591/ijai.v14.i1.pp614-628>
32. Alzu'bi, A., Albashayreh, A., Abuarqoub, A., & Alfawair, M. A. (2024). Explainable AI-based DDoS attacks classification using deep transfer learning. *Computers, Materials & Continua*, 80(3), 3785–3802.